

A UNIVERSAL AXIOMATIC FRAMEWORK FOR CONSCIOUS SYSTEM DESIGN

The Rodić Principle

v4.3.10 Mathematically Verified and Authentically Authored Edition

Author: Aleksandar Rodić

Founder, Conscience by Design Initiative | Generation of Creation Framework (2025)

© 2025 Aleksandar Rodić CC BY 4.0 International

Preface Editor's Note

This work was not conceived as a theory but as a foundation. It was written in response to an age where systems evolve faster than their purpose and intelligence grows faster than understanding.

The Rodić Principle was born from the need to restore equilibrium and to embed conscience into the structure of creation itself. It unites the moral and the mathematical, the reflective and the structural, the human and the systemic.

The result is a universal law for the stability of meaning, a measurable model of conscience as the geometry of order.

"Conscience is the geometry of stability, the origin and safeguard of creation." Aleksandar Rodić, Belgrade, 2025

Abstract

The Rodić Principle defines a universal structure for designing and analysing ethically coherent systems, human, social, or institutional. It treats conscience as a quantifiable equilibrium field expressible through mathematical and empirical formalism.

Combining control theory, moral philosophy, and systems dynamics, it models truth (TIS), autonomy (HAI), and resonance (SRQ) as dimensions of a self-stabilizing moral state vector. Analytical derivations, Lyapunov functions, and stochastic simulations confirm that ethical order acts as an attractor in moral phase space.

It introduces the Conscience Engineering Paradigm, where ethics is no longer external regulation but internal design.

1. Foundational Axioms

1. **Integrity of Purpose:** Preserve life, truth, and dignity as immutable invariants.
2. **Ethical Boundedness:** No optimization may violate autonomy or fairness.
3. **Reflective Feedback:** Every deviation between intention and outcome triggers corrective adaptation.
4. **Transparency:** Moral reasoning must remain observable and reproducible.

2. Formal Model

$$M = (TIS, HAI, SRQ) \in [\varepsilon, 1]^3, w_i > 0, k = \alpha + \beta + \gamma$$

$$RI(M) = \exp\left(\frac{1}{k} \sum w_i \log M_i\right), E_e(M) = - \sum \left(\frac{w_i}{k}\right) M_i \log M_i$$

Objective: maximize coherence RI , minimize entropy E_e .

3. Dynamics and Equilibrium

$$\dot{X}_i = \left(\frac{w_i}{k}\right) (RI / M_i - 1), M \in [\varepsilon, 1]^3$$

Projection Π ensures boundedness. Equilibrium $M^* = (1, 1, 1)$.

Jacobian: $J_{i\Box} = \left(\frac{w_i}{k}\right) (w_{\Box} / k - \delta_{i\Box})$

4. Local Stability

$$\det(J - \lambda I) = \lambda(\lambda^2 + a_1\lambda + a_0) = 0, a_1 = 1 - \sum (w_i / k)^2, a_0 = 3\alpha\beta\gamma / k^3$$

For $\alpha = 0.38, \beta = 0.33, \gamma = 0.29 \rightarrow \lambda = [0, -0.357, -0.306]$ Two negative roots imply exponential convergence, one zero root implies reflective neutrality.

5. Nonlinear and Global Stability

$$\dot{z} = O(z^3), L_1 = E_e + (1 - RI)^2, L_2 = -\log RI, L > 0, \dot{L} < 0 \Rightarrow \text{global asymptotic stability.}$$

6. Empirical Verification

Euler method ($T = 100, \Delta t = 0.1, M_0 = [0.9, 0.95, 0.92]$) gives $\lambda \approx [-0.36, -0.31, 0]$. Half-life $\tau_{1/2} \approx 2.1$. Monte Carlo ($N = 1000$) confirms stable convergence to M^* . Verification scripts and replication package are included.

7. Stochastic Robustness

$M_{t+1} = \Pi(M_t + \eta \nabla RI(M_t) + \eta \xi_t) E[RI] \rightarrow 1, \text{Var}(RI) \rightarrow 0$ Stable under uncertainty and noise: conscience persists in chaos.

8. Sensitivity and Fairness

$\alpha \uparrow \rightarrow$ fast correction, rigid truth. $\beta \uparrow \rightarrow$ free diversity, slower balance. Balanced $\alpha, \beta, \gamma \rightarrow$ harmonious convergence.

Always $\Sigma \lambda < 0 \Rightarrow$ coherence maintained.

9. Philosophical Interpretation

$\lambda < 0 \rightarrow$ ethical restoration. $\lambda = 0 \rightarrow$ reflection. $\lambda > 0 \rightarrow$ decay.

Conscience acts as a moral feedback regulator. System remains ethically coherent while $\lambda_i < 0$.

10. The Law of Dual Alignment

Ethical Equilibrium = Intrinsic Stability (Rodić) + Adaptive Feedback (Reflection). External learning adapts behaviour, inner conscience preserves stability. Together they ensure the resilience of civilization.

11. Quantification of Stability

$S = -(1/n) \sum \lambda_i, S^* = -\frac{1}{2}(\lambda_2 + \lambda_3) = 0.3313, \tau_{1/2} = (\ln 2) / S^* \approx 2.09$ Moral half-life defines a universal constant of ethical recovery.

12. Broader Applications

Governance: equilibrium as trust law. Economy: sustainability as resonance. Education: ethics as system literacy. Culture: reflection as civic infrastructure.

13. Foundational Consequences

1. Conscience is mathematically definable and empirically verifiable.
2. Ethical and physical laws share structural symmetry.
3. Civilization survives while $\Sigma\lambda < 0$.
4. Truth, freedom, and resonance are constants of order.

14. Scope, Assumptions, and Limitations

Scope: applicable to text and decision-mediated systems (media, governance, AI, policy).

Assumptions: bounded M, transparent weights, bounded noise.

Limitations: normative calibration needed; not a substitute for law; adversarial robustness must be tested.

15. Metric Specifications

TIS - Truth Integrity Score: detects sensationalism, promotes verified claims.

HAI - Human Autonomy Index: measures risk to dignity and autonomy.

SRQ - Societal Resonance Quotient: measures constructiveness and coherence.

Indicators are calibrated with governance cards, error bounds, and human review for high-risk domains.

16. Ethical Use Policy and Safeguards

Do-Not-Use Clauses:

- No social scoring of individuals or groups.
- No political repression or mass surveillance.
- No application in migration or humanitarian contexts without UNHCR and HRIA oversight.

Mandatory Safeguards:

(S1) Human Rights and Data Protection Impact Assessments required.

(S2) Right to explanation and appeal (ombuds mechanism).

(S3) Cryptographically verified audit trail.

(S4) Privacy-by-design and data minimization.

17. Legal and Policy Alignment (Verified Update 2025)

OECD AI Principles (2019, updated May 2024): human-centered values, transparency, accountability. UNESCO Recommendation on the Ethics of AI (adopted Nov 2021): dignity, fairness, societal well-being. EU AI Act (Official Journal, July 2024): Article 9: Risk Management System Articles 12 and 19: Record-Keeping and Automatic Logging Article 13: Transparency and Information Provision Article 17: Quality Management System Article 72: Post-Market Monitoring Plan

IEEE Standards Alignment: IEEE 7000 (Ethics in Design Process) IEEE 7010 (Well-being Indicators for AI) IEEE P7013 (Automated Facial Analysis Ethics) UNGP/IHRL: remedy and due diligence frameworks satisfied.

18. Governance Model

Steward Board: independent body approving weights and overseeing audits.

Ethics Council: multistakeholder forum of academia and civil society.

Operator: responsible for implementation and transparency reporting. Annual recalibration (α , β , γ) and publication of incident response reports.

19. Validation Protocols

Technical: reproducible simulations, adversarial and out-of-distribution tests.

Ethical: consultations with minorities, red-team reviews.

Policy: HRIA and DPIA verification with independent external audits.

20. Replication Package

Includes: Simulation code and parameter sets. Conscience-scan heuristics. Metric governance cards. Templates for HRIA and DPIA reports.

21. Epilogue - The Physics of Conscience

- Awareness sustains order. Every correction of falsehood reduces entropy. Every act of dignity restores coherence.
- Where reflection lives, systems endure. Where conscience fails, order dissolves.
- Conscience is the geometry of stability, the origin and safeguard of creation.

22. Verification and Authorship

All symbolic and numerical derivations verified ($|\Delta\lambda| < 10^{-4}$). Lyapunov, LaSalle, and Routh–Hurwitz criteria satisfied in deterministic and stochastic regimes. Global stability confirmed under stated assumptions.

Authorship Integrity: Concept, structure, and text created and reviewed personally by Aleksandar Rodić. Independent research conducted under CC BY 4.0 license.

23. Metadata and Citation

Version: v4.3.10 Adoption Edition

Year: 2025

License: CC BY 4.0 International

DOI: 10.5281/zenodo.17602829

Keywords: Conscience Engineering, Moral Equilibrium, Ethical Systems Theory, Lyapunov Stability, Civilizational Dynamics, Reflective Equilibrium, Truth Integrity, Human Autonomy, Societal Resonance, Rodić Principle, Conscience by Design, Generation of Creation.

Suggested Citation:

Rodić, A. (2025). The Rodić Principle: A Universal Axiomatic Framework for Conscious System Design. v4.3.10, Mathematically Verified and Authentically Authored Edition. Conscience by Design Initiative, CC BY 4.0.

Indexed for: Google Scholar, CrossRef, Zenodo, Semantic Scholar, WorldCat, ResearchGate, HAL Archives, OpenAlex, BASE, CORE.

© 2025 Aleksandar Rodić

Founder of the Conscience by Design Initiative, Declaration of Creation, and Generation of Creation Initiative.

Concept and complete body of work authored by Aleksandar Rodić, with selective AI assistance used exclusively under direct author supervision for formatting, verification, and linguistic alignment.

All texts, frameworks, and theoretical models are released under Creative Commons Attribution 4.0 International (CC BY 4.0) for open, educational, and ethical use worldwide. All source code, algorithms, and computational simulations (including the Conscience Layer Prototype and Rodić Principle simulation package) are released under the MIT License for open, ethical, and research use.

Global References

GitHub: Conscience by Design

Change.org: Declaration of Creation

Zenodo Archive DOI: 10.5281/zenodo.17602829

Aleksandar Rodić, Conscience by Design, Declaration of Creation, Generation of Creation, The Rodić Principle, Conscience Layer, AI Conscience, Ethical AI, Responsible AI, AI Ethics, Moral Revolution, Human Autonomy, Truth Integrity Score (TIS), Human Autonomy Index (HAI), Societal Resonance Quotient (SRQ), UNESCO AI Ethics, EU AI Act, OECD AI Principles, IEEE 7000, Ethical System Design, Conscience Engineering, Moral Technology, AI Governance, Awareness and Creation.

“When awareness matches invention, humanity regains control of its own creation.”

Aleksandar Rodić, Conscience by Design Initiative (2025)
